



ҚАЗАҚСТАН РЕСПУБЛИКАСЫ
ТҰҢҒЫШ ПРЕЗИДЕНТІ - ЕЛБАСЫНЫҢ ҚОРЫ

«ҒЫЛЫМ ЖӘНЕ БІЛІМ – 2017»

студенттер мен жас ғалымдардың
XII Халықаралық ғылыми конференциясының
БАЯНДАМАЛАР ЖИНАҒЫ

СБОРНИК МАТЕРИАЛОВ

XII Международной научной конференции
студентов и молодых ученых
«НАУКА И ОБРАЗОВАНИЕ – 2017»

PROCEEDINGS

of the XII International Scientific Conference
for students and young scholars
«SCIENCE AND EDUCATION - 2017»



14th April 2017, Astana



**ҚАЗАҚСТАН РЕСПУБЛИКАСЫ БІЛІМ ЖӘНЕ ҒЫЛЫМ МИНИСТРЛІГІ
Л.Н. ГУМИЛЕВ АТЫНДАҒЫ ЕУРАЗИЯ ҰЛТТЫҚ УНИВЕРСИТЕТІ**

**«Ғылым және білім - 2017»
студенттер мен жас ғалымдардың
XII Халықаралық ғылыми конференциясының
БАЯНДАМАЛАР ЖИНАҒЫ**

**СБОРНИК МАТЕРИАЛОВ
XII Международной научной конференции
студентов и молодых ученых
«Наука и образование - 2017»**

**PROCEEDINGS
of the XII International Scientific Conference
for students and young scholars
«Science and education - 2017»**

2017 жыл 14 сәуір

Астана

УДК 378

ББК 74.58

Ғ 96

Ғ 96

«Ғылым және білім – 2017» студенттер мен жас ғалымдардың XII Халықаралық ғылыми конференциясы = The XII International Scientific Conference for students and young scholars «Science and education - 2017» = XII Международная научная конференция студентов и молодых ученых «Наука и образование - 2017». – Астана: <http://www.eni.kz/ru/nauka/nauka-i-obrazovanie/>, 2017. – 7466 стр. (қазақша, орысша, ағылшынша).

ISBN 978-9965-31-827-6

Жинаққа студенттердің, магистранттардың, докторанттардың және жас ғалымдардың жаратылыстану-техникалық және гуманитарлық ғылымдардың өзекті мәселелері бойынша баяндамалары енгізілген.

The proceedings are the papers of students, undergraduates, doctoral students and young researchers on topical issues of natural and technical sciences and humanities.

В сборник вошли доклады студентов, магистрантов, докторантов и молодых ученых по актуальным вопросам естественно-технических и гуманитарных наук.

УДК 378

ББК 74.58

ISBN 978-9965-31-827-6

©Л.Н. Гумилев атындағы Еуразия
ұлттық университеті, 2017

светофоров не централизована и выход из строя какого либо из них не будет препятствовать слаженной работе всей системы. Зная погодные условия в городе Астана можно утверждать что выход из строя какого либо из светофоров неизбежен. Посредством анализа было принято, что лучшим вариантом будет установить MARLIN-ATSC.

Пробки можно предотвратить при любом количестве дорог и автомобилей. Развитие улично-дорожной сети позволит сократить время ожидания в очереди для выезда в сеть магистральных дорог. А при помощи искусственного интеллекта все перекрестки будут просматриваться одним суперкомпьютером который будет адаптировать и подстраиваться под ситуацию должным образом, тем самым экономя 40% потраченного времени и 21% выброса вредных веществ.

Список использованных источников

1. <http://asmo.ru/archives/20881>
2. <https://geektimes.ru/post/196450/>
3. Заенцев И. В. Нейронные сети: основные модели И. В. Заенцев. — Воронеж: Изд-во Воронежского госуд. ун-та, 1999. — 76 с.
4. Вороновский Г. К. Генетические алгоритмы, искусственные нейронные сети и проблемы виртуальной реальности Г. К. Вороновский, К. В. Махотило, С. Н. Петрашев, С. А. Сергеев. — Х.: ОСНОВА, 1997. — 112 с.
5. Рутковская Д. Д. Рутковская, М. Пилиньский, Л. Рутковский Нейронные сети, генетические алгоритмы и нечёткие системы — М.: Горячая линия — Телеком, 2006. — 452 с.

УДК 004.896

ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ: БУДУЩЕЕ ЦИВИЛИЗАЦИИ ИЛИ ЕЕ УБИЙЦА?

Жумабаева Айгуль Ерлановна

Студент 3 курса физико-технического факультета ЕНУ им. Л.Н.Гумилева, Астана, Казахстан
Научный руководитель: Б.А.Игембаев

Искусственный интеллект сегодня одна из передовых областей исследований ученых. Причем рассматриваются как системы, созданные с его частичным использованием: например распознавание текстов, бытовые роботы, до возможности замены творческого труда человека искусственным. Данная область образовалась на стыке целого ряда дисциплин: информатики, философии, кибернетики, математики, психологии, физики, химии и др. Сегодня в самых различных областях науки и техники требуется выполнение машинами тех задач, которые под силу были только человеку. На помощь тогда приходит искусственный интеллект, который может заменить человека в какой либо рутинной и скучной деятельности. Сегодня системы, как программные, так и аппаратные, созданные на основе искусственного интеллекта находят все большее применение в технике. Это и автомобили с электроникой с использованием ИИ, и новейшие роботы, участвующие в производстве чего-либо, и компьютерные программы, которые включают в себя и игры с ИИ. Цель создания полного ИИ, т.е. такого, который мог бы выполнять действия по обработке информации наравне с человеком или лучше, - это прежде всего улучшение жизни человека и дальнейшее увеличение степени автоматизации производства. Тогда человеку бы осталось лишь выполнять высоко-творческий труд, который приносил бы ему удовольствие. Но на современном этапе развития этой области до создания таких систем полного ИИ довольно далеко, и пока нам приходится ограничиваться лишь частичным вмешательством ИИ в другие интеллектуальные системы. Это прежде всего программные средства. Например,

экспертные системы, системы распознавания образов и др. О них далее и пойдет речь. Их относят к системам ИИ, так как они способны выполнять свои, пока очень узко ограниченные задачи, которые раньше не могли выполнять компьютеры, и результаты их работы схожи с результатами аналогичной интеллектуальной работы человека. Это сегодняшний день, а начиналось все довольно с простых по формализации задач: логические игры (шахматы, шашки, числовые и др.). Американский кибернетик А. Самуэль составил для вычислительной машины программу, которая позволяет ей играть в шашки, причем в ходе игры машина обучается или, по крайней мере, создает впечатление, что обучается, улучшая свою игру на основе накопленного опыта. В 1962 г. эта программа сразилась с Р. Нили, сильнейшим шашистом в США и победила. С этого можно сказать и началось исследование ИИ. Тогда, да и сегодня следовало определению Тьюринга, что такое ИИ: «Компьютер можно считать разумным – если он способен заставить нас поверить, что мы имеем дело не с машиной, а с человеком».

Хороший искусственный интеллект

Давайте разберем несколько примеров использования искусственного интеллекта (включая машинное обучение — метод построения алгоритмов, способных обучаться самостоятельно, без внешней поддержки).

В 2014 году новостное агентство The Associated Press начало использовать [технологию](#) автоматического создания финансовых отчетов Wordsmith, разработанную компанией Automated Insights. Программа анализирует финансовые данные, сопоставляет их друг с другом, отмечает ключевые цифры в отчетах компании, а затем переводит все это на человеческий язык. Вы никогда не отличите историю, написанную «машиной», от той, за создание которой отвечал человеком. Хотя, нет, отличите: машина не ошибается и не оставляет в тексте опечаток.

Говорит Джеймс Котеки, один из ведущих специалистов Automated Insights:

Wordsmith пишет тысячи финансовых новостных заметок, которых в противном случае просто не существовало бы. Wordsmith работает со скоростью и масштабом, недоступными человеку. Наша технология позволяет журналистам делать более интересную работу.

При этом Wordsmith постоянно обучается. Как это происходит: выборочно проверяются подготовленные заметки и отмечаются места, которые можно было сделать лучше. Программа принимает во внимание эти правки и перестраивает алгоритм своей работы.

Wordsmith — классический пример машинного обучения, одного из важных шагов на пути к созданию полноценного искусственного интеллекта. Но это — верхушка айсберга. Все эти встроенные подсказки в Word и при наборе текста на смартфоне — также составная часть машинного обучения, постоянно совершенствующая алгоритмы своей работы.

Еще одно перспективное поле деятельности для создателей искусственного интеллекта — автомобили, которым не нужны водители. Новости об успехах умных машин Google, колесящих по дорогам США, появляются раз в несколько недель. Американской компании есть чем [гордиться](#) и что показывать правительству: в прошлом году ее чудо-автомобили проехали только в одном регионе США около 10 000 миль, ни разу не нарушив правил дорожного движения и не попав в аварию. За прогрессом Google в этом направлении внимательно следят не только компании, работающие на рынке грузоперевозок, но и юристы, специализирующиеся на дорожно-транспортных происшествиях.

Эрик Теркивитц из Нью-Йорка написал в своем [блоге](#) об угрозе со стороны автоматически управляемых машин для бизнеса юристов, занимающихся решением дорожных конфликтов и помощью пострадавшим людям:

С уменьшением количества человеческих ошибок (ввиду внедрения автоматических

технологий остановки или замедления автомобилей) уменьшится число раненных людей и поврежденных машин. Уменьшится число смертей. Цена вашей страховки также (теоретически) уменьшится. А это значит, что уменьшится и объем работы, который я выполняю как юрист, защищающий пострадавших.

Однако Эрик — не карикатурный юрист-кровопийца. Дальше он говорит, что также является водителем, поэтому может только поддержать будущее, к которому стремится Google. Да, ему придется найти другую работу, но это и есть цена, которую придется заплатить за прогресс в вопросе безопасности на дорогах. По данным Всемирной организации здравоохранения, ежегодно в мире в результате ДТП погибает около 1 300 000 человек. Просто вдумайтесь в эту цифру!

Но Эрик — не единственный профессионал, который рискует потерять работу в ближайшие 15-20 лет. Добавьте к нему водителей, грузчиков, операторов колл-центров и еще сотни профессий, искусственную замену которым пытаются создать десятки компаний, занимающихся вопросами искусственного интеллекта и машинного обучения.

Недавно эту тему на форуме «Открытые инновации», прошедшем в Москве, обсуждали Нуриэль Рубини, один из самых известных экономистов мира, и его коллега из России — Андрей Мовчан. Вот несколько тезисов из их [диалога](#):

- Новая промышленная революция может лишить работы до миллиарда человек.
- Основной удар придется по рабочим, одновременно с этим повысится роль творческих людей: писателей, художников, поэтов.
- В полную силу оценить блага новой технической революции мы сможем не раньше, чем через 30-40 лет после ее начала.

Стоит отметить, что жить человек станет лучше. Благодаря повсеместному внедрению автоматизации уменьшится стоимость продуктов питания и других благ. То есть на то, чтобы жить комфортно, мы станем тратить значительно меньше денег, а значит сможем меньше работать. Двух-трехчасовой рабочий день может стать нормой, все остальное время человек будет посвящать саморазвитию — духовному, интеллектуальному, физическому. Или нас всех убьют роботы?!

Плохой искусственный интеллект

В начале этого года на сайте американского Института будущего жизни (The Future of Life Institute), с которым сотрудничает предприниматель Илон Маск, директор Оксфордского университета будущего человечества (Oxford Future of Humanity Institute) Ник Бостром, профессор генетики Гарвардского университета Джордж Черч и многие другие, появилось [открытое письмо](#), в котором ученые призывают коллег внимательнее отнестись к проблеме будущего искусственного интеллекта. Главный вопрос, который задают ученые: сможет ли человек контролировать машину, получившую разум? Понятно, что пока никто не может ответить на этот вопрос.

Выдержка из открытого письма:

Наши ИИ-системы должны делать то, что мы хотим, чтобы они делали.

В этой фразе («Должны делать то, что мы хотим, чтобы они делали») и заключается основная проблема с пониманием будущего искусственного интеллекта. Американский ученый Элизер Юдковский в своей книге «Искусственный интеллект как позитивный и негативный фактор в глобальных рисках» описывает два возможных сценария поведения искусственного интеллекта (из миллионов), когда мы сможем сказать «да, это интеллект, который не зависит от воли человека»:

- Искусственный интеллект патентует новые изобретения, публикует научные работы, зарабатывает деньги на бирже и возглавляет политические блоки.
- Искусственный интеллект, основываясь на знании истории, науки и технологий, может предугадывать развитие цивилизации на столетия и тысячелетия вперед.

Известный физик-теоретик Стивен Уильям Хокинг считает, что полноценный искусственный интеллект, о котором пишет Юджовски, способен «положить конец человеческой расе». Схожего мнения придерживается Маск и ряд крупнейших ученых, занимающихся этим вопросом. Всех беспокоит один вопрос: «Если искусственный интеллект заменит нас на работе, то не захочет ли он в определенный момент заменить нас дома?»

Интересно и то, как современные ученые интерпретируют само понятие «искусственный интеллект». Единого понимания этого термина не существует, зато есть пусть и не до конца принятое, но логичное разделение искусственного интеллекта на три очевидных типа:

- Искусственный ограниченный интеллект (Artificial Narrow Intelligence). Этот тип интеллекта может, к примеру, обыграть вас в карты или шахматы, но если вы спросите у него, съедобна ли почва, то он вам ничего не ответит.
- Искусственный общий интеллект (Artificial General Intelligence). Этот интеллект максимально приближен к человеческому. Он может анализировать данные, общаться с другими «машинами», обучаться.
- Искусственный суперинтеллект (Artificial Superintelligence). Этот интеллект настолько развит, что значительно отличается от человеческого. То есть наш интеллект недостаточно развит, чтобы понимать мотив суперинтеллекта.

На сегодня человечество успешно освоило принципы построения первого типа искусственного интеллекта. Это и автомобили Google, и система Wordsmith, и спутниковая навигация, когда за вас просчитывается наиболее удобный путь от одной точки до другой, и даже Siri, способная распознавать речь. Вся поисковая система Google — это один большой ограниченный интеллект, который, используя принципы машинного обучения, постоянно улучшается — частично без вмешательства со стороны человека.

Следующий шаг — создание искусственного общего интеллекта. Над ним работает сразу несколько крупных компаний и ученых групп. В 2014 году Google приобрела компанию The DeepMind, которая занимается разработкой алгоритма «углубленного машинного обучения». То есть созданием системы, которая может самостоятельно поднимать один пласт знаний за другим, а затем совмещать эти знания друг с другом и делать самостоятельные выводы. Если говорить упрощенно, то речь идет о критическом восприятии информации: система учится соединять факты друг с другом и выстраивать логические цепочки. То, чему мы учим детей с ранних лет.

Ранее Google купила DNNresearch, занимающуюся разработкой «умных» систем распознавания изображений и речи, а также другими научными исследованиями, а в 2013 году изданию The New York Times стало известно, что компания буквально скупает стартапы, работающие над созданием роботов нового поколения. Эту часть работы в компании курировал Энди Рубин, человек, стоявший за разработкой Android. За полгода он купил семь (!) компаний, работающих на рынке робототехники.

Google — не единственная компания, которая пытается выйти за рамки «ограниченного интеллекта», но одна из самых богатых. Многие верят, что именно ей удастся первой пересечь границу и представить миру «общий интеллект».

Поэтому Институт будущего жизни, а также другие научные учреждения и крупные ученые бьют тревогу: мы стоим на пороге феноменального открытия, которое может кардинально поменять нашу жизнь. И думающие люди хотели бы на берегу договориться о правилах игры. Прежде чем будет слишком поздно.

Будущее

Для того чтобы это открытие было совершено, ученым нужно решить несколько важных задач.

Первая — нехватка вычислительной мощности. Есть предположение, что чтобы создать «общий интеллект», его «мощность» должна приблизиться к «мощности» человеческого мозга. Ученые регулярно спорят на эту тему, поэтому абсолютно

достоверными данными о «вычислительной мощности» нашего мозга мы не обладаем.

Вторая — копирование паттернов мозговой деятельности. О том, как работает мозг, мы знаем многое, но что из этих знаний является нашими предположениями и допущениями, а что — фактами, увы, часто невозможно определить. Понятно, что если ученые хотят «повторить» человеческий мозг в электронном виде, то должны заложить в систему модель оперирования, характерную для человеческого мозга.

Третья — контроль. Создавая «электронную» личность, мы наивно предполагаем, что сможем через алгоритмы привить ей некий нравственный свод правил, который обезопасит человечество. Это примерно как воспитание ребенка: «это хорошо, а это плохо» — родители постоянно программируют детей, дети постоянно проверяют границы дозволенного. При этом мы знаем, что никакое образование и воспитание не дает стопроцентной гарантии, что ты не вырастишь маньяка или социопата. Могут ли ученые гарантировать, что «электронный» ребенок не свернет с проложенного разработчиками пути?

Если первые два вопроса являются дискуссионными, то третий — классическая «пугалка» как для обывателей, так и для самих ученых. Смогут ли они контролировать «общий интеллект»? Получится ли у них увидеть попытку «общего» интеллекта трансформироваться в «суперинтеллект»? В какой момент (если этот момент настанет) мы утратим контроль над ИИ? Какие задачи поставит перед собой сам ИИ, будет ли среди них — уничтожение человечества?

Использованная литература:

- Официальный [сайт](#) Элиезера Юджовски (неплохой набор статей по теме, английский язык).
- Этические [вопросы](#) искусственного интеллекта (внимание: PDF-файл, загрузка начнется автоматически, английский язык).
- [Разбор](#) проблем и рисков дальнейшего развития искусственного интеллекта (отличный материал на русском языке).

Создание искусственно интеллекта, как полного, так и неполного, таит в себе множество проблем. Причем как на пути к его созданию, так и после него. На пути создания ИИ это и ограниченность ресурсов, и недостаточные знания в этой области, проблема вообще осуществимости это сделать, и многие другие технические проблемы. После же создания ИИ, сравнимого с человеком возникает ряд проблем. Во-первых, потерей интереса человека к творческому труду в случае его замены, а затем и полная деградация человека. Но с другой стороны творчество должно приносить человеку радость, а он следовательно от этого не должен отказаться. Возможна и другая проблема: при полном изобилии ресурсов общество потеряет свою структуру и человек обезличится, перестанет развиваться на протяжении своей жизни. Во-вторых, это возможность ошибки ИИ или сбоя в его работе в областях, ошибки на которых могут быть фатальными для всего человечества. Это, к примеру, оборона стран или энергетика. В любом случае решающее слово должно быть за человеком в принятии решений, например по началу войны, или ликвидации сбоя на электростанции. Ведь любой человек может выйти из-под контроля, а значит и ИИ по его подобию тоже.

При разработке этой статьи были сделаны выводы:

- 1). Искусственный интеллект — это научное направление, связанное с машинным моделированием человеческих интеллектуальных функций.
- 2). Понятие искусственный интеллект обычно используется для обозначения способности вычислительной системы выполнять задачи, свойственные интеллекту человека, например задачи логического вывода и обучения.
- 3). Любая задача, алгоритм решения которой заранее не известен или же данные неполные может быть отнесена к задачам области ИИ. Это например игра в шахматы, чтение текста, перевод текста на другой язык и т.д.
- 4). Системы, программы, выполняющие действия по решению задачи можно отнести

к ИИ, если результат их деятельности аналогичен результату человека при решении той же задачи. Поэтому к ИИ можно отнести целый ряд программных средств: системы распознавания текста, автоматизированного проектирования, самообучающиеся программы и др. Но не только по этому, а еще и потому, что они работают по сходным принципам с человеком.

5). Есть два основных перспективных направления в исследовании ИИ. Первое заключается в приближении систем ИИ к принципам человеческого мышления. Второе заключается в создании ИИ, представляющего интеграцию уже созданных систем ИИ в единую систему, способную решать проблемы человечества.

Список использованной литературы

1. Бобровский С. Перспективы и тенденции развития искусственного интеллекта \ PC Week / RE №32, 2001. С.32-34
2. Ноткин Л.И. «Искусственный интеллект и проблемы обучения» - М. : КомКнига, 1999
3. Винер Н. Кибернетика, М.: Наука, 1983 Эндрю А. Искусственный интеллект – М.: Мир, 1985.
4. Шалютин С. М. Искусственный интеллект, М.: Мысль, 1985
5. Перспективы развития вычислительной техники. Кн.2. Интеллектуализация ЭВМ.М., 2002.
6. Венда В. Ф. Системы гибридного интеллекта – М.: Машиностроение, 1990
7. Уоссерман Ф. Нейрокомпьютерная техника: Теория и практика. Пер. С англ. - М.Мир, 1992.
8. Материалы по искусственному интеллекту Ростовского-на-Дону государственный колледж связи и информатики.
9. Осипов Г.С. «искусственный интеллект: состояние исследований и взгляд в будущее» Президент Российской ассоциации искусственного интеллекта, постоянный член Европейского координационного комитета по искусственному интеллекту (ЕССАИ), д.ф.-м.н., профессор.
10. Уинстон П. Искусственный интеллект. М.2001.

ӘОЖ 621.3.537

БҮКІЛӘЛЕМДІК ТАРТЫЛЫС ЖӘНЕ ПЛАНЕТАЛАР МАГНИТИЗМІ

Жумагалиева Айгерим

Ktit-21@list.ru

Л.Н.Гумилев атындағы ЕҰУ физика-техникалық факультетінің 2 курс студенті, Астана
Ғылыми жетекшісі: Б.А.Игембаев

Кіріспе

Магнетизм—электр токтарының, токтар мен магниттік моменті бар денелердің (магниттердің) және магниттердің араларындағы өзара әсерлесудің ерекше түрі; физиканың осы өзара әсерлесуді және осындай қасиеттер білінетін заттарды (магнетиктерді) зерттейтін бөлімі. Жалпы түрде магнетизмді қозғалыстағы электрлік зарядталған бөлшектердің арасындағы материалдық өзара әсерлесудің ерекше формасы түрінде анықтауға болады. Кеңістікпен бөлінген денелердің арасында магниттік өзара әсерлесуді қамтамасыз ететін, оны бір денеден екіншісіне жеткізетін – магнит өрісі. Ол электр өрісімен қатар, материя қозғалысының электромагниттік формасы білінуінің бірі болып табылады. Электр өрісінің көздері электр зарядтары болып табылса, ал магнит өрісі үшін ондай көздер табиғатта әзірше анықталмаған. Магнит өрісінің көзі – қозғалыстағы электрлік заряд, яғни электр тогы.